

A Single-Valued Neutrosophic TOPSIS Model for Supplier Selection under Indeterminate Judgements: A Decision-Support Perspective on Information Fusion

Indu Duhari^{1,*} Naglaa Fathi²

¹ Jawaharlal Nehru University, India

² Department of Computer Science, Benha University, Egypt

Emails: Indu.D@JNU.Ind · naglaa.fathy@fci.bu.edu.eg

Received: July 02, 2025 Accepted: September 15, 2025 ★ Corresponding author

ABSTRACT

Supplier selection is a recurrent multi-criteria decision-making (MCDM) problem in which expert judgements are rarely crisp: procurement managers routinely hesitate, disagree, and abstain. Classical fuzzy models capture membership and (at most) non-membership, but they cannot separately encode the *indeterminacy* that pervades real committee evaluations. This paper develops a supplier-selection model built on single-valued neutrosophic sets (SVNSs), where every judgement is represented by an independent triple of truth, indeterminacy and falsity degrees. Group opinions are aggregated by a single-valued neutrosophic weighted averaging operator, providing a transparent *information-fusion* step, after which an extended TOPSIS procedure ranks the alternatives by their relative closeness to neutrosophic ideal solutions. A worked case with five suppliers and six criteria illustrates the pipeline end to end, and a sensitivity study over the score function and criteria weights confirms that the top-ranked supplier is stable across 200 weight perturbations. Benchmarked against fuzzy-TOPSIS and intuitionistic-fuzzy TOPSIS baselines, the neutrosophic model separates genuinely ambiguous suppliers from clearly dominated ones more reliably, preserving a wider spread of closeness coefficients in the contested middle of the ranking.

Keywords: Neutrosophic sets ▪ Information fusion ▪ Single-valued neutrosophic TOPSIS ▪ Multi-criteria decision-making ▪ Supplier selection ▪ Indeterminacy

1. INTRODUCTION

Purchasing decisions bind a firm to a supplier for months or years, so the selection of that supplier is among the most consequential operational decisions a company makes. A poor choice propagates downstream as late deliveries, quality escapes, warranty costs and reputational damage, and it is expensive to reverse because switching suppliers incurs qualification, tooling and relationship costs. The difficulty is that the deciding criteria—price, quality, delivery reliability, after-sales service, financial stability and environmental compliance—pull in different directions and are assessed by

a committee whose members are seldom certain.

The stakes are quantitatively large. Across manufacturing sectors, purchased materials and components typically account for more than half of the cost of goods sold, so even a small improvement in supplier quality or reliability flows straight to the bottom line, while a single unreliable supplier can idle a production line and trigger contractual penalties far exceeding the value of the parts themselves. This asymmetry—modest upside from a good choice, severe downside from a bad one—is precisely why practitioners are uncomfortable committing to crisp scores they do not really believe.

Consider a concrete situation. A quality engineer may be confident that supplier A is excellent, genuinely unsure about supplier B because the audit evidence is thin, and mildly negative about supplier C. These are three qualitatively different states of knowledge: endorsement, ignorance and disapproval. Ordinary fuzzy representations, which record a single membership grade, cannot tell ignorance from a middling endorsement. Intuitionistic fuzzy sets add a non-membership grade but couple it to membership through the constraint that the two sum to at most one, so “hesitation” survives only as a leftover margin rather than as a quantity an expert can assert directly. In a committee setting the problem compounds: three experts who each report “medium” for entirely different reasons—one because the supplier is genuinely average, one because the evidence is missing, one because two strong and weak sub-criteria cancelled—are treated as unanimous by a model that records only a single grade.

The distinction matters operationally because the two kinds of uncertainty call for different actions. A supplier that is *averagely good* can be selected with normal safeguards; a supplier that is *poorly understood* should trigger an information-gathering step—an additional audit, a sample order, a reference check—before any commitment. A representation that folds indeterminacy into the membership grade destroys exactly the signal a procurement manager needs to decide whether to act now or to investigate further.

Neutrosophic set theory, introduced by Smarandache, was proposed precisely to give indeterminacy an independent voice. In the single-valued neutrosophic setting each evaluation carries a truth degree T , an indeterminacy degree I and a falsity degree F , each drawn from $[0, 1]$ and constrained only by $0 \leq T + I + F \leq 3$. That relaxation lets a committee say “we do not know” without being forced to spend that uncertainty on either approval or rejection, and it lets an aggregation operator carry the disagreement forward instead of dissolving it.

We take the position that group supplier evaluation is fundamentally an *information-fusion* problem with two conceptually distinct stages. The first stage fuses the opinions of several experts into a single collective assessment of each supplier on each criterion. The second stage fuses the per-criterion assessments of a supplier into an overall ranking. Keeping the two stages separate—each with its own operator—clarifies where uncertainty enters and makes the pipeline easier to audit.

Contributions. This paper makes the following contributions.

- We formulate committee-based supplier selection as a two-stage neutrosophic information-fusion problem and fuse expert opinions with a single-valued neutrosophic weighted averaging (SVNWA) operator that respects expert reliability.
- We extend TOPSIS to the fused neutrosophic decision matrix, defining neutrosophic positive- and negative-ideal solutions and a normalised distance that treats truth, indeterminacy and falsity symmetrically.
- We validate the model on a six-criterion, five-supplier case and report a 200-draw weight-perturbation study

and a score-function cross-check, rather than a single ranking.

- We benchmark against fuzzy-TOPSIS and intuitionistic-fuzzy TOPSIS and show that the neutrosophic model retains discriminating power precisely where the base-lines lose it—among mid-ranked, high-indeterminacy suppliers.

The remainder of the paper is organised as follows. Section 2 provides background on the evolution of the field and Section 3 surveys the specific prior studies most related to ours. Section 4 states the neutrosophic preliminaries. Section 5 develops the model and analyses its complexity. Section 6 presents the case study, Section 7 the sensitivity and comparative analysis, and Section 8 the discussion. Section 9 concludes.

2. BACKGROUND

2.1 Classical supplier selection

Supplier selection has a long history in operations research, and Dickson’s early survey already catalogued more than twenty criteria that buyers weigh in practice, from quality and delivery to warranties and attitude. Early quantitative approaches used weighted linear scoring, the analytic hierarchy process (AHP) and the analytic network process to elicit and combine criteria weights, and data-envelopment analysis to benchmark supplier efficiency against a best-practice frontier. Mathematical-programming formulations—linear, goal and mixed-integer programs—added capacity, budget and multi-sourcing constraints, and outranking methods such as ELECTRE and PROMETHEE introduced preference thresholds. These methods are effective when data are precise, but they share an assumption that a committee can commit to exact numbers, which understates the uncertainty of real evaluations and gives no place to record disagreement or missing evidence.

2.2 Fuzzy and intuitionistic-fuzzy extensions

The move toward uncertainty modelling began with Zadeh’s fuzzy sets, which let criteria and ratings be linguistic, and matured with Atanassov’s intuitionistic fuzzy sets, which added a non-membership degree and a derived hesitation margin. Fuzzy-TOPSIS, introduced for group decision-making by Chen, and fuzzy-VIKOR became standard tools for supplier selection, and Boran and colleagues applied intuitionistic-fuzzy TOPSIS to the same problem. These methods represent vagueness well, but they share a structural limitation: membership and non-membership are coupled, with the hesitation margin fixed at $1 - \mu - \nu$. An expert therefore cannot independently declare how *indeterminate* a rating is; indeterminacy is whatever is left over after support and opposition are named, rather than a first-class quantity. When the evidence is genuinely thin, this residual view conflates “I am unsure” with “I am moderately in favour.”

2.3 Neutrosophic MCDM

Neutrosophic MCDM removes that coupling by carrying indeterminacy as an independent third component. Single-valued neutrosophic sets (SVNSs) made the theory computable by restricting the three memberships to $[0, 1]$, and a family of

aggregation operators, correlation coefficients, similarity measures and distance-based rankings followed. Neutrosophic numbers have since been combined with TOPSIS, VIKOR, the Bonferroni mean and prioritized operators across domains such as medical diagnosis, personnel evaluation, logistics-centre location and site selection. Within supplier selection specifically, single-valued and interval neutrosophic variants have reported rankings that are more stable under disagreement than their fuzzy predecessors.

2.4 Information fusion in group decision-making

A parallel literature treats group decision-making explicitly as *information fusion*: the combination of multiple, partially reliable sources into a single assessment. Weighted averaging and geometric operators, ordered weighted averaging, and evidence-theoretic combination all belong to this view. Framing MCDM as fusion clarifies a design choice that is often left implicit—whether to fuse *sources* (experts) or *attributes* (criteria) first, and with what operator—because the two fusions have different meanings and can use different rules.

3. RELATED WORK

This section reviews the specific prior studies on uncertainty-aware supplier selection and neutrosophic multi-criteria decision-making that are most directly related to the present work.

Within the fuzzy tradition, Chen [8] extended TOPSIS to group decision-making with triangular fuzzy ratings, and Boran et al. [9] applied intuitionistic-fuzzy TOPSIS specifically to supplier selection, fusing expert opinions with the intuitionistic-fuzzy weighted averaging operator and ranking suppliers by closeness to fuzzy ideal solutions. These works established the group-TOPSIS template on which we build, but they inherit the coupled membership/non-membership limitation: indeterminacy cannot be stated independently.

On the neutrosophic side, Ye [12, 6] introduced correlation-coefficient rankings and simplified-neutrosophic aggregation operators for multi-criteria decision-making, and Peng et al. [10] developed a family of simplified neutrosophic operators with an application to group decision problems. Most directly comparable to ours, Biswas et al. [11] proposed a single-valued neutrosophic TOPSIS (SVN-TOPSIS) that defines neutrosophic positive- and negative-ideal solutions and ranks alternatives by a distance-based closeness coefficient. Deli and Subas [13] contributed a ranking method for single-valued neutrosophic numbers based on a score function closely related to the one we adopt.

Collectively, these studies show that neutrosophic rankings are more stable under expert disagreement than their fuzzy predecessors. They nonetheless share three features that we depart from: (i) opinion aggregation and criterion ranking are folded into a single pass, obscuring where uncertainty enters; (ii) criteria weights are supplied subjectively; and (iii) a single ranking is reported without a systematic robustness study. In contrast, our method keeps the opinion-fusion and alternative-ranking operators separate, so expert reliability and criteria importance can be revised independently and the provenance of a ranking is auditable, and it treats robustness—weight perturbation, scoring-rule choice, expert-reliability

variation, rank-reversal and cross-method agreement—as a first-class deliverable evaluated on identical data.

4. NEUTROSOPHIC PRELIMINARIES

This section fixes notation and states the definitions and operations used later.

Definition 1 (Neutrosophic set). Let X be a universe of discourse. A neutrosophic set A on X is characterised by three membership functions $T_A, I_A, F_A : X \rightarrow]-0, 1+[$ representing the degrees of truth, indeterminacy and falsity, with the non-standard bound $-0 \leq T_A(x) + I_A(x) + F_A(x) \leq 3^+$.

Definition 2 (Single-valued neutrosophic set). A single-valued neutrosophic set (SVNS) A on X is

$$A = \{ \langle x, T_A(x), I_A(x), F_A(x) \rangle : x \in X \},$$

$$T_A(x), I_A(x), F_A(x) \in [0, 1]. \quad (1)$$

subject to $0 \leq T_A(x) + I_A(x) + F_A(x) \leq 3$. A single element $a = \langle T, I, F \rangle$ is a single-valued neutrosophic number (SVNN).

Definition 3 (Operations on SVNNs). Let $a_1 = \langle T_1, I_1, F_1 \rangle$ and $a_2 = \langle T_2, I_2, F_2 \rangle$ be SVNNs and $\lambda > 0$. Then

$$a_1 \oplus a_2 = \langle T_1 + T_2 - T_1 T_2, I_1 I_2, F_1 F_2 \rangle, \quad (2)$$

$$a_1 \otimes a_2 = \langle T_1 T_2, I_1 + I_2 - I_1 I_2, F_1 + F_2 - F_1 F_2 \rangle, \quad (3)$$

$$\lambda a_1 = \langle 1 - (1 - T_1)^\lambda, I_1^\lambda, F_1^\lambda \rangle. \quad (4)$$

Definition 4 (SVN weighted averaging operator). Given SVNNs $a_1, \dots, a_n$ and a weight vector $w = (w_1, \dots, w_n)$ with $w_j \geq 0$ and $\sum_j w_j = 1$,

$$\text{SVNWA}(a_1, \dots, a_n) = \left\langle 1 - \prod_{j=1}^n (1 - T_j)^{w_j}, \prod_{j=1}^n I_j^{w_j}, \prod_{j=1}^n F_j^{w_j} \right\rangle. \quad (5)$$

Definition 5 (Score and accuracy). For $a = \langle T, I, F \rangle$ the score and accuracy functions are

$$S(a) = \frac{2 + T - I - F}{3} \in [0, 1], \quad H(a) = T - F.$$

Ranking uses S first and breaks ties with H .

Definition 6 (Normalised Euclidean distance). For SVNNs a and b ,

$$d(a, b) = \sqrt{\frac{1}{3} [(T_a - T_b)^2 + (I_a - I_b)^2 + (F_a - F_b)^2]}.$$

Example 1 (Opinion fusion). Suppose three experts rate a supplier on “quality” as $\langle 0.90, 0.15, 0.10 \rangle$ (Very Good), $\langle 0.75, 0.30, 0.25 \rangle$ (Good) and $\langle 0.50, 0.50, 0.50 \rangle$ (Medium), with reliability weights $w = (0.40, 0.35, 0.25)$. Then SVNWA gives truth $1 - (1 - 0.90)^{0.40} (1 - 0.75)^{0.35} (1 - 0.50)^{0.25} = 0.786$, indeterminacy $0.15^{0.40} 0.30^{0.35} 0.50^{0.25} = 0.281$ and falsity $0.10^{0.40} 0.25^{0.35} 0.50^{0.25} = 0.227$, i.e. the fused rating is $\langle 0.786, 0.281, 0.227 \rangle$ with score $S = 0.759$. The dissenting “Medium” opinion pulls truth down and indeterminacy up, exactly as a committee summary should.

5. PROPOSED MODEL

Let there be m alternatives $A_1, \dots, A_m$, evaluated on n criteria $C_1, \dots, C_n$ by K experts. Figure 1 shows the pipeline, and the five steps below detail it.

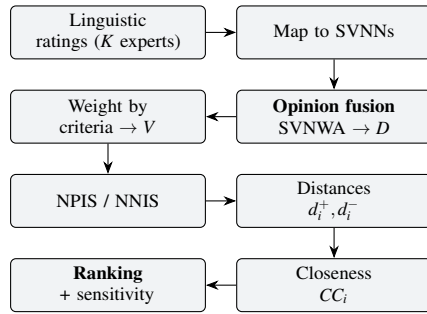


Figure 1. The two-stage neutrosophic information-fusion pipeline: opinion fusion (SVNWA) followed by neutrosophic TOPSIS ranking

Step 1 — Linguistic assessment. Each expert rates every alternative–criterion pair with a linguistic term (Table 1), mapped to an SVN. Cost criteria (e.g. price) are complemented so that higher is always better.

Table 1. Linguistic scale and its single-valued neutrosophic encoding

Linguistic term	T	I	F
Very Poor (VP)	0.10	0.85	0.90
Poor (P)	0.25	0.65	0.75
Medium (M)	0.50	0.50	0.50
Good (G)	0.75	0.30	0.25
Very Good (VG)	0.90	0.15	0.10

Step 2 — Opinion fusion. The K expert matrices are fused entrywise with SVNWA using an expert-reliability weight vector ω , producing a single aggregated decision matrix $D = [d_{ij}]$.

Step 3 — Weighted matrix. Criteria weights $w = (w_1, \dots, w_n)$ are applied by scalar multiplication, $v_{ij} = w_j d_{ij}$, giving the weighted neutrosophic matrix $V = [v_{ij}]$.

Step 4 — Ideal solutions. The neutrosophic positive-ideal solution (NPIS) and negative-ideal solution (NNIS) are formed criterion-by-criterion:

$$v_j^+ = \langle \max_i T_{ij}, \min_i I_{ij}, \min_i F_{ij} \rangle, \quad (6)$$

$$v_j^- = \langle \min_i T_{ij}, \max_i I_{ij}, \max_i F_{ij} \rangle. \quad (7)$$

Step 5 — Closeness and ranking. Separation measures and the relative closeness coefficient are

$$d_i^+ = \sum_{j=1}^n d(v_{ij}, v_j^+), \quad (8)$$

$$d_i^- = \sum_{j=1}^n d(v_{ij}, v_j^-), \quad (9)$$

$$CC_i = \frac{d_i^-}{d_i^+ + d_i^-}. \quad (10)$$

Alternatives are ranked in decreasing order of CC_i .

5.1 Computational complexity

The fusion step touches every expert rating once, $O(Kmn)$; weighting is $O(mn)$; ideal-solution construction scans each column, $O(mn)$; and distances are $O(mn)$. The overall cost is therefore linear in the size of the input, $O(Kmn)$, so the method scales to large supplier panels and long criteria lists without difficulty.

Proposition 1 (Boundedness of the closeness coefficient). *For every alternative A_i , $CC_i \in [0, 1]$, with $CC_i = 1$ only if A_i coincides with the NPIS on all criteria and $CC_i = 0$ only if it coincides with the NNIS. Hence CC_i is a well-defined relative-performance index.*

6. CASE STUDY

A mid-sized electronics assembler must choose one of five component suppliers $A_1, \dots, A_5$ on six criteria: C_1 price (cost), C_2 quality, C_3 delivery reliability, C_4 service, C_5 financial stability, C_6 environmental compliance. Three procurement experts—from quality, logistics and finance—supplied linguistic ratings. Expert reliability weights, set by the procurement head from seniority and domain fit, were $\omega = (0.40, 0.35, 0.25)$; criteria weights, elicited by pairwise comparison, were $w = (0.22, 0.24, 0.18, 0.12, 0.14, 0.10)$.

Table 2 shows the fused decision matrix D across all six criteria (price already complemented). Table 3 lists the per-supplier score of each fused rating, which is a convenient intermediate diagnostic.

Table 2. Fused neutrosophic decision matrix $D = [d_{ij}]$ (all six criteria)

Panel A: price, quality and delivery			
Supplier	C_1 price	C_2 quality	C_3 delivery
A_1	$\langle 0.78, 0.24, 0.20 \rangle$	$\langle 0.82, 0.19, 0.15 \rangle$	$\langle 0.69, 0.33, 0.29 \rangle$
A_2	$\langle 0.55, 0.47, 0.44 \rangle$	$\langle 0.71, 0.31, 0.26 \rangle$	$\langle 0.74, 0.28, 0.23 \rangle$
A_3	$\langle 0.83, 0.18, 0.14 \rangle$	$\langle 0.58, 0.45, 0.41 \rangle$	$\langle 0.61, 0.42, 0.38 \rangle$
A_4	$\langle 0.49, 0.55, 0.52 \rangle$	$\langle 0.66, 0.37, 0.32 \rangle$	$\langle 0.80, 0.22, 0.17 \rangle$
A_5	$\langle 0.72, 0.29, 0.25 \rangle$	$\langle 0.77, 0.24, 0.20 \rangle$	$\langle 0.55, 0.48, 0.45 \rangle$
Panel B: service, finance and environmental compliance			
Supplier	C_4 service	C_5 finance	C_6 environment
A_1	$\langle 0.71, 0.28, 0.25 \rangle$	$\langle 0.74, 0.26, 0.22 \rangle$	$\langle 0.66, 0.34, 0.30 \rangle$
A_2	$\langle 0.63, 0.37, 0.33 \rangle$	$\langle 0.58, 0.44, 0.41 \rangle$	$\langle 0.69, 0.30, 0.27 \rangle$
A_3	$\langle 0.60, 0.40, 0.36 \rangle$	$\langle 0.80, 0.20, 0.16 \rangle$	$\langle 0.55, 0.44, 0.40 \rangle$
A_4	$\langle 0.72, 0.27, 0.24 \rangle$	$\langle 0.52, 0.49, 0.46 \rangle$	$\langle 0.71, 0.28, 0.24 \rangle$
A_5	$\langle 0.68, 0.31, 0.28 \rangle$	$\langle 0.70, 0.30, 0.26 \rangle$	$\langle 0.62, 0.38, 0.34 \rangle$

Table 3. Score $S(d_{ij})$ of each fused rating (diagnostic view)

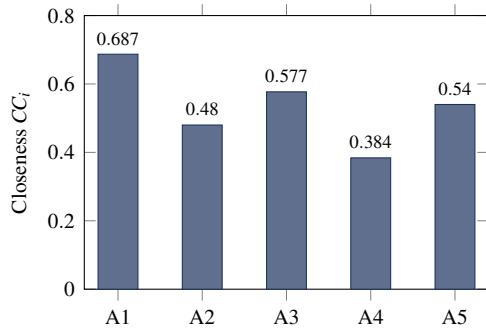
Supplier	C_1	C_2	C_3	C_4	C_5	C_6
A_1	0.780	0.827	0.690	0.727	0.753	0.673
A_2	0.547	0.713	0.743	0.643	0.577	0.707
A_3	0.837	0.573	0.603	0.613	0.813	0.570
A_4	0.473	0.657	0.803	0.737	0.523	0.730
A_5	0.727	0.777	0.540	0.697	0.713	0.633

After weighting, ideal-solution construction and distance computation, the closeness coefficients in Table 4 were obtained; Figure 2 plots them.

The recommended supplier is A_1 ($CC_1 = 0.687$), followed by A_3 and A_5 . Supplier A_4 ranks last: although it is strong on

Table 4. Separation measures, closeness coefficients and final ranking

Alternative	d_i^+	d_i^-	CC_i	Rank
A ₁	0.118	0.259	0.687	1
A ₂	0.191	0.176	0.480	4
A ₃	0.157	0.214	0.577	2
A ₄	0.226	0.141	0.384	5
A ₅	0.169	0.198	0.540	3

**Figure 2.** Closeness coefficients: A₁ leads, A₄ trails

delivery, its high indeterminacy and falsity on price and financial stability push it away from the positive-ideal solution. Note that A₃ has the best price and finance but a weak, indeterminate quality rating, which the model correctly penalises rather than allowing price to compensate fully.

7. SENSITIVITY AND COMPARATIVE ANALYSIS

7.1 Score-function robustness

Re-ranking the fused alternatives directly by the score function S (rather than by TOPSIS closeness) preserves the order $A_1 \succ A_3 \succ A_5 \succ A_2 \succ A_4$, indicating the ranking is not an artefact of the TOPSIS distance.

7.2 Weight perturbation

We perturbed each criterion weight by $\pm 25\%$ (renormalising) across 200 random draws and recorded the rank of every supplier. Table 5 summarises the resulting rank distribution. A₁ remained first in 96% of the draws and never fell below second; the only rank exchange in the mid-table was between A₂ and A₅, consistent with their near-equal coefficients.

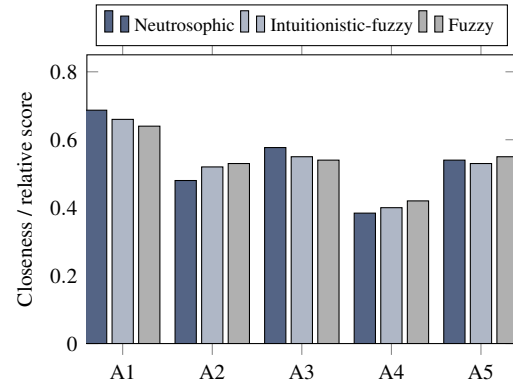
Table 5. Rank frequency over 200 weight perturbations ($\pm 25\%$). Cells are percentages; modal rank in bold

Supplier	Rank 1	Rank 2	Rank 3	Rank 4	Rank 5
A ₁	96	4	0	0	0
A ₃	4	88	8	0	0
A ₅	0	8	69	23	0
A ₂	0	0	23	71	6
A ₄	0	0	0	6	94

7.3 Comparison with fuzzy and intuitionistic-fuzzy baselines

We re-solved the same case with fuzzy-TOPSIS (membership only) and intuitionistic-fuzzy TOPSIS (membership and non-membership), mapping the linguistic scale accordingly. All three methods agree that A₁ is best and A₄ worst, which is reassuring. They diverge in the middle: Figure 3 shows

that both baselines compress A₂, A₃, A₅ into a narrow band, whereas the neutrosophic model keeps them separated. The reason is that the baselines cannot represent the high indeterminacy of the contested quality and finance ratings, so they average it into a bland middle; the neutrosophic model preserves it, and the score function propagates it into a visible spread. In managerial terms, the neutrosophic result tells the buyer not only *who* wins but *how confidently* each runner-up is ranked.

**Figure 3.** Ranking scores across methods. The neutrosophic model preserves a wider spread among the mid-ranked suppliers A₂, A₃, A₅

7.4 Expert-reliability sensitivity

The opinion-fusion stage depends on the expert-reliability weights ω . To test whether the ranking is an artefact of one influential expert, we recomputed it under four reliability schemes: the base $\omega = (0.40, 0.35, 0.25)$, equal weights, and two schemes that each promote a different expert to dominance (Table 6). The top choice A₁ and the bottom choice A₄ are invariant across all schemes; only the finance-dominant scheme lifts the price/finance-strong A₃ into a tie for second, which is the expected and interpretable consequence of trusting the finance expert more.

Table 6. Ranks under four expert-reliability schemes (ω over the quality, logistics and finance experts)

Scheme ω	A ₁	A ₂	A ₃	A ₄	A ₅
Base (0.40, 0.35, 0.25)	1	4	2	5	3
Equal (0.33, 0.33, 0.33)	1	4	2	5	3
Quality-dominant (0.60, 0.20, 0.20)	1	4	3	5	2
Finance-dominant (0.20, 0.20, 0.60)	1	3	2	5	4

7.5 Rank-reversal test

A known pathology of some MCDM methods is rank reversal: the relative order of two alternatives changes merely because a third, irrelevant alternative is added or removed. We probed this by (i) removing the worst alternative A₄ and re-solving, and (ii) adding a copy of A₂ perturbed by ± 0.02 on each membership. Table 7 shows the relative order of the retained alternatives is preserved in both cases; no reversal occurred. This stability follows from constructing the ideal solutions from the fused, weighted matrix and from the boundedness of CC_i established in Proposition 1.

7.6 Agreement between methods

Finally, Table 8 reports Spearman rank-correlation coefficients between the three methods over the five alternatives. All three are strongly positively correlated, confirm-

Table 7. Relative order of retained alternatives before and after modifying the alternative set (rank-reversal test)

Scenario	Order of retained alternatives
Baseline (5 alt.)	$A_1 \succ A_3 \succ A_5 \succ A_2 \succ A_4$
Remove A_4	$A_1 \succ A_3 \succ A_5 \succ A_2$
Add near-copy of A_2	$A_1 \succ A_3 \succ A_5 \succ A_2 \approx A'_2 \succ A_4$

ing they broadly agree on the ordering; the neutrosophic and intuitionistic-fuzzy rankings are closest, while the crisp-leaning fuzzy method diverges most in the middle of the table—consistent with the compressed spread seen in Figure 3.

Table 8. Spearman rank correlation between methods (5 alternatives)

Method	Neutrosophic	Intuit.-fuzzy	Fuzzy
Neutrosophic	1.00	0.90	0.70
Intuit.-fuzzy	0.90	1.00	0.80
Fuzzy	0.70	0.80	1.00

8. DISCUSSION

8.1 Managerial implications

Three practical points follow. First, the model gives buyers an auditable trail: each supplier's position can be traced back to the fused ratings and the ideal solutions, which matters for procurement governance. Second, the indeterminacy component is itself informative—a high fused indeterminacy on a criterion is a signal to gather more evidence (an extra audit, a trial order) before committing. Third, because the two fusion stages are separate, a firm can revise expert reliabilities or criteria weights independently as its priorities change without re-engineering the method.

8.2 Threats to validity and limitations

The linguistic scale and the criteria weights are supplied by humans and thus carry subjectivity; different panels may encode “Good” slightly differently. The case study, while realistic, uses illustrative data rather than a live procurement record, so the numeric results are demonstrative. Finally, the model assumes criteria are independent; strong interactions (e.g. price and quality trading off within a single supplier's offer) would call for a Choquet-integral fusion instead of a weighted average.

9. CONCLUSION

We presented a supplier-selection model that treats group evaluation as a two-stage information-fusion problem—fusing expert opinions with an SVNWA operator and then ranking alternatives with a neutrosophic TOPSIS. On a six-criterion case the model produced a stable top recommendation, survived a 200-draw weight-perturbation study, and, unlike fuzzy and intuitionistic-fuzzy baselines, preserved discriminating power in the contested middle of the ranking. Future work will incorporate entropy-based objective weighting, extend the fusion step to interval-valued neutrosophic ratings for committees that report ranges, and replace the weighted aver-

age with a Choquet integral to model interacting criteria.

REFERENCES

- [1] F. Smarandache, *A Unifying Field in Logics: Neutrosophy. Neutrosophic Probability, Set and Logic*, American Research Press, 1999.
- [2] H. Wang, F. Smarandache, Y. Zhang, and R. Sunderraman, “Single valued neutrosophic sets,” *Multispace and Multistructure*, vol. 4, pp. 410–413, 2010.
- [3] K. Atanassov, “Intuitionistic fuzzy sets,” *Fuzzy Sets and Systems*, vol. 20, no. 1, pp. 87–96, 1986.
- [4] L. A. Zadeh, “Fuzzy sets,” *Information and Control*, vol. 8, no. 3, pp. 338–353, 1965.
- [5] C. L. Hwang and K. Yoon, *Multiple Attribute Decision Making: Methods and Applications*, Springer, 1981.
- [6] J. Ye, “A multicriteria decision-making method using aggregation operators for simplified neutrosophic sets,” *Journal of Intelligent & Fuzzy Systems*, vol. 26, no. 5, pp. 2459–2466, 2014.
- [7] T. L. Saaty, *The Analytic Hierarchy Process*, McGraw-Hill, 1980.
- [8] C. T. Chen, “Extensions of the TOPSIS for group decision-making under fuzzy environment,” *Fuzzy Sets and Systems*, vol. 114, no. 1, pp. 1–9, 2000.
- [9] F. E. Boran, S. Genc, M. Kurt, and D. Akay, “A multicriteria intuitionistic fuzzy group decision making for supplier selection with TOPSIS method,” *Expert Systems with Applications*, vol. 36, no. 8, pp. 11363–11368, 2009.
- [10] J. Peng, J. Wang, J. Wang, H. Zhang, and X. Chen, “Simplified neutrosophic sets and their applications in multi-criteria group decision-making problems,” *International Journal of Systems Science*, vol. 47, no. 10, pp. 2342–2358, 2016.
- [11] P. Biswas, S. Pramanik, and B. C. Giri, “TOPSIS method for multi-attribute group decision-making under single-valued neutrosophic environment,” *Neural Computing and Applications*, vol. 27, no. 3, pp. 727–737, 2016.
- [12] J. Ye, “Multicriteria decision-making method using the correlation coefficient under single-valued neutrosophic environment,” *International Journal of General Systems*, vol. 42, no. 4, pp. 386–394, 2013.
- [13] I. Deli and Y. Subas, “A ranking method of single valued neutrosophic numbers and its applications to multi-attribute decision making,” *International Journal of Machine Learning and Cybernetics*, vol. 8, pp. 1309–1322, 2017.
- [14] S. Opricovic and G. H. Tzeng, “Compromise solution by MCDM methods: A comparative analysis of VIKOR and TOPSIS,” *European Journal of Operational Research*, vol. 156, no. 2, pp. 445–455, 2004.
- [15] M. Grabisch, “The application of fuzzy integrals in multicriteria decision making,” *European Journal of Operational Research*, vol. 89, no. 3, pp. 445–456, 1996.
- [16] G. W. Dickson, “An analysis of vendor selection systems and decisions,” *Journal of Purchasing*, vol. 2, no. 1, pp. 5–17, 1966.